

Electronic Supplementary Information ESM_1.pdf

Estimation of uncertainty from duplicate measurements: new quantification procedure in the case of concentration-dependent precision

Václav Synek, Sylvie Kříženecká

Faculty of Environment, Jan Evangelista Purkyně University in Ústí nad Labem, Pasteurova 3632/15, CZ-400 96
Ústí nad Labem, Czech Republic

Correspondence to Sylvie Kříženecká

E-mail address: sylvie.krizenecka@ujep.cz

Text S1	Application of standard statistical procedures.....	2
Text S1.1	Assessing correlation.....	2
Text S1.2	Checking the homogeneity and normality of the estimate sets.....	2
Text S1.3	The diagrams of $ d_i $ or $ d_{ri} $ against $\bar{c}_i$	2
Text S2	The chosen concentration distributions of the simulated duplicate datasets	3
Fig. S1		3
Table S1.		5
Text S3	Results of estimates obtained from the datasets with $n = 20\,000$ and their discussion	6
Text S3.1	Series of the sequential estimates	6
Table S2.		6
Text S3.2	The variability, bias, and errors of the estimates	7
Table S3.		8
Table S4.		9
Text S4	Overview of the estimates obtained from datasets with $n = 200$ and some their statistics	10
Table S5.		11
Table S6.		12
Fig. S2.		13
Table S7.		14
Table S8.		15
Text S5	Correction proportion from the sum of the squared differences.....	16
Table S9.		16
Table S10.		17
Table S11.		18
Text S6	Statistical significance of the s_0^2 and s_r^2 estimates obtained by regression procedures.....	20
Table S12		20
Text S7	Correlation between the parameter estimates obtained by all used procedures	21
Text S8	References.....	21
Table S13.		22
Table S14.		23

The list of symbols and abbreviations used is given in the published article.

Text S1 Application of standard statistical procedures

Text S1.1 Assessing correlation

Since some of the processed datasets were probably distributed non-normally or outliers could have occurred in them, Spearman's rank correlation coefficient was used. The coefficients were calculated using software TIBCO Statistica version 13.3.

For the parameter estimates from the datasets with $n = 20\,000$, a correlation between the RME values of the s_0 and s_r final estimates was investigated. These estimates are summarized in Table S4 in ESM_1.pdf, there are 20 pairs of the s_0 and s_r estimates. The found correlation coefficient is given and discussed in Text S3.2. The critical value of the correlation coefficient for $n = 20$ and at $\alpha = 0.05$ is 0.447 [26]. Furthermore, four matrices of correlation coefficients were calculated from the parameter estimates obtained by processing the datasets with $n = 200$ distributed *uniformly* and *exponentially* (see Tables S13 and S14). These coefficients assess the relationships between the s_0 or s_r estimates obtained by pairs of different estimating procedures from the ten datasets with a given distribution type (see Tables S5 and S6 in ESM_1.pdf for these estimates). Each matrix summarises correlation coefficients for all possible pairs of the estimating procedures used and their variants, the critical value at $\alpha = 0.05$ is 0.648 [26]. Correlation coefficients were also used to investigate the relationships between the correction proportions and the errors of parametric estimates or differences of parallel pairs of estimates or RSD of estimate sets (for the sources data see Tables S3, S4, S9, S10 and S11, for evaluation of the correlation coefficients see Text S5 in ESM_1.pdf).

Text S1.2 Checking the homogeneity and normality of the estimate sets

The distributions of the sets with tens of s_0 or s_r estimates, which were obtained using the different procedures from the datasets with $n = 200$ (for these estimates see Tables S5 and S6), were investigated using box plots created by software TIBCO Statistica. These plots (see Fig. S2 in ESM_1.pdf) identify mild and extreme outliers according to the inner and outer fences and allow visual comparisons of variability and skewness of the investigated sets. Furthermore, the Grubbs tests for homogeneity were applied to test minimum and maximum values and pairs of outliers at opposite ends of the sets [4]. The normality of the sets was tested by the Kolmogorov-Smirnov/Lilliefors test [12] and [1s - Text s8 in ESM_1.pdf]. The skewness and kurtosis coefficients of the sets were also investigated [2s - Text s8 in ESM_1.pdf]. The statistics found are summarised in Tables S7 and S8, evaluation see Text S4 in ESM_1.pdf.

Text S1.3 The diagrams of $|d_i|$ or $|d_{ri}|$ against $\bar{c}_i$

These diagrams are placed in Figures 3 and 4. The s_c values were calculated according to Eq. (8) by substituting the values of s_0 and s_r : (i) the true values; (ii) the estimates obtained by the RMS WLS procedure as the best estimation procedure; (iii) the estimates obtained by a variant of the proposed procedure. The diagrams present primarily two cases of such datasets, where the values of the estimates obtained both by the regression and by the proposed procedure were very close to the true values (one dataset with a uniform distribution of concentrations and the other with an exponential distribution). In these examples, it can be assumed that the arrangement of the plotted points will be representative for both tested types of the concentration distribution under the given true s_0 and s_r values. The other diagrams depict selected cases of a few datasets with interesting divergences from the characteristic arrangements.

Text S2 The chosen concentration distributions of the simulated duplicate datasets

The simulated datasets for estimating the s_0 and s_r parameters should provide sufficient information on both variance components. This means that the range of the simulated concentrations should encompass both the range with sufficiently low concentrations, where s_0^2 is the dominant component of the s_c^2 variance, and the range with sufficiently high concentrations, where $s_r^2 c^2$ is the dominant component of s_c^2 . A sufficiently dominant portion is certainly 95% or more, which means that at $s_0 = 0.15$ cu and $s_r = 0.07$, s_0^2 and $s_r^2 c^2$ are clearly dominant, when $c < 0.49$ cu (this is near LOD) and $c > 9.3$ cu, respectively.

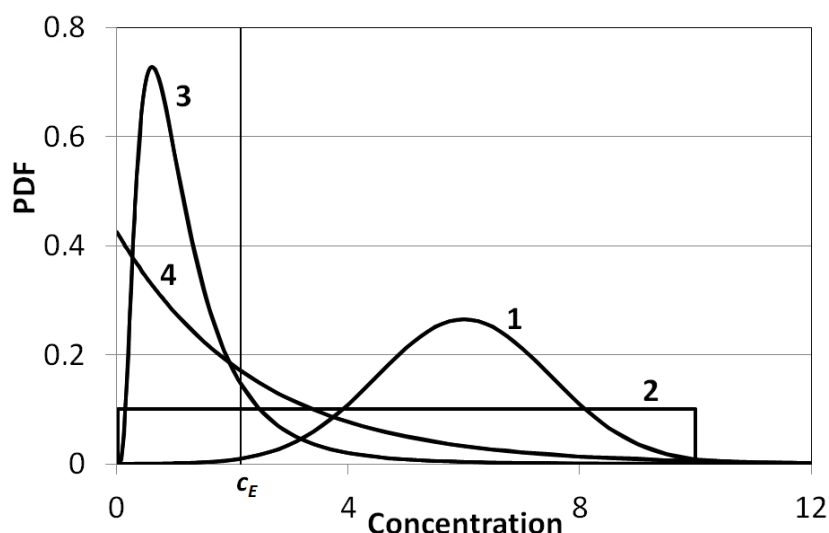


Fig. S1 Selected Probability Density Functions (PDF) of the analyte concentration used in the simulation of the datasets; distribution (1) normal, (2) uniform, (3) log-normal and (4) exponential; $c_E = s_0/s_r$

The sets of duplicates were simulated with different concentration distributions (see Fig. S1 in ESM_1.pdf), since the relative frequency of the concentrations in various concentration ranges could affect the estimates obtained by the proposed procedure. As it is possible to expect that Equations (14) and (15) will work properly, if the concentrations of the samples monitored cover the whole concentration range uniformly, first of all, a uniform - rectangular distribution was selected, although the concentration distribution of laboratory samples will probably be different. Then the most commonly assumed type of distribution was chosen, i.e., a normal distribution. Log-normal and exponential distribution models were also tested, Their choice is justified in the following paragraph.

The concentration distribution of duplicate results accumulated within IQC will follow the distribution of concentrations of samples coming into the laboratory. In the case of trace analysis, the predominant samples usually have concentrations at the bottom of the measurement interval, working range, of the measurement procedure. The distribution is limited from below by zero, i.e., values lower than LOD. The most frequent values are those corresponding to the usual level of the analyte content in the samples of a given matrix, e.g., in environmental samples from a given area according to the usual content of the natural substance or pollutant being analysed. Samples with higher concentrations of the analyte occur at a low frequency and samples with extremely high concentrations are exceptional, i.e., the distribution of concentrations is skewed to higher values. This concept corresponds to the log-normal distribution model. If the range of usual analyte concentrations were around and below the LOD, this would be more of an exponential distribution. For example, when checking whether the

concentration of an analyte in samples does not exceed a hygienic limit which is higher than the LOQ, then a large portion of the measured samples can have the analyte concentration at or below the LOD.

Table S1 (ESM_1.pdf) summarises the parameters of the concentration populations with the four chosen distributions, statistics of the samples of simulated true concentrations, μ_i , and statistics of the means, $\bar{c}_i$, of simulated concentrations measured in duplicate (the statistics only for the cases of the 1st series of the simulated datasets with $n = 20\,000$). It can be seen that the statistics of the μ_i samples do not practically differ from the corresponding theoretical parameters of those chosen probability distributions. The minimum values of these samples approached zero from the right. The maximum values varied from 10 cu to 23 cu, see the uniform and exponential distributions, respectively. The statistics computed from the $\bar{c}_i$ samples were close to the statistics computed from the samples of the true concentrations, because the values μ_i and $\bar{c}_i$ differed only by the added random errors. These differences between μ_i and $\bar{c}_i$ values resulted in small proportions of negative values in the $\bar{c}_i$ datasets. The highest proportion of the negative $\bar{c}_i$ values appeared in the dataset with the exponential distribution; only the dataset with the normal distribution was without any negative value.

The table also shows the proportions of values of concentration lower than the equivalence concentration, i.e., the proportions of duplicate differences suitable for estimating s_0 . The proportions of differences suitable for estimating s_r represent the complements of 100%. It can be seen that in the case of the chosen normal and log-normal distributions, the datasets of the duplicates have negligible and very small proportions of differences suitable for estimating s_0 and s_r , respectively.

In Fig. S1 (ESM_1.pdf) it can be seen the concentrations with maxima of PDFs, modes. For the normal and log-normal distributions these concentrations are 6 and 0.613 cu. For the exponential distribution, PDF reaches its maximum when the concentration approaches zero, for the uniform distribution, PDF does not have a maximum.

Table S1. Parameters of the concentration populations with the chosen distribution types and statistics computed from the samples with 20 000 simulated values of the true concentration, μ_i , and from the samples of means, $\bar{c}_i$, of the simulated duplicates; only the 1st series of the simulated datasets was processed. Symbols μ , σ , a , b , λ - parameters of the applied distributions; $c_E = s_0/s_r = 2.14$; concentration values are assumed to be expressed in a concentration unit

Distribution type	Selected parameters	Studied set	Mean	Median	SD	Minimum	Maximum	% of neg. values	% of values $\leq c_E$
Normal $N(\mu, \sigma^2)$	$\mu = 6$ $\sigma = 1.5$	population	6	6	1.5	$\rightarrow -\infty$	$\rightarrow \infty$	0.003	0.506
		μ_i sample	6.006	6.003	1.492	0.1159	12.422	0	0.460
		$\bar{c}_i$ sample	6.009	6.002	1.528	0.2908	13.313	0	0.500
Uniform $U(a, b)$	$a = 0$ $b = 10$	population	5	5	2.896	0	10	0	21.4
		μ_i sample	5.012	5.007	2.884	0.0009	9.9999	0	21.3
		$\bar{c}_i$ sample	5.012	5.013	2.901	-0,3367	11,557	0.42	21.4
Log-normal $LN(\mu, \sigma^2)$	$\mu = 0$ $\sigma = 0.7$	population	1.278	1	1.016	$\rightarrow 0$	$\rightarrow \infty$	0	86.2
		μ_i sample	1.282	1.003	1,005	0.0672	14.368	0	85.8
		$\bar{c}_i$ sample	1,282	1.007	1.014	-0,2121	13.197	0.09	85.6
Exponential $Exp(\lambda)$	$\lambda = 1/2.35$ $= 0.4255$	population	2.35	1.629	2.35	$\rightarrow 0$	$\rightarrow \infty$	0	59.8
		μ_i sample	2.371	1.656	2.353	7.04E-06	23.429	0	59.3
		$\bar{c}_i$ sample	2.371	1.652	2.363	-0.2678	24.639	1.78	59.4

Text S3 Results of estimates obtained from the datasets with $n = 20\,000$ and their discussion

Text S3.1 Series of the sequential estimates

Table S2 (ESM_1.pdf) summarises the series of the sequential s_0 and s_r estimates obtained by the chosen variants of the proposed procedure from the 1st series of the datasets with 20 000 duplicates simulated with the chosen types of concentration distribution. The first estimate in each column is the zeroth estimate calculated using Eq. (3) and the last is the final estimate, i.e., the estimate of s_0 and s_r considered to be stabilised; when the difference between two subsequent parameter estimates was less than 0.5%, the parameter with a larger difference is considered. At least 4 iterations were performed in each sequence.

Table S2. Series of the sequential estimates of s_0 and s_r obtained by chosen variants of the proposed procedure from the 1st series of the datasets with 20 000 duplicates simulated with the chosen types of the concentration distribution (Dis.); j – number of iterations; n_0 and n_r – numbers of duplicates for calculation of s_0 and s_r ; Dif. – difference between the two last estimates in the given column

Dis.	Normal						Exponential			
n_0	10 000	4 000	2 600	400	10 000					
n_r	10 000	16 000	17 400	19 600	10 000					
Estimates expressed in % of the true values										
j	s_0	s_r	s_0	s_r	s_0	s_r	s_0	s_r	s_0	s_r
0	252.6	103.7	212.0	106.0	195.2	106.7	150.5	107.6	107.2	124.9
1	196.3	69.9	157.4	77.2	139.9	81.1	109.6	90.5	99.9	95.8
2	162.7	84.9	129.4	91.3	114.0	94.4	99.7	98.9	99.2	100.1
3	144.0	91.2	116.4	96.3	103.2	98.7	97.6	100.4	99.2	100.5
4	134.2	94.1	110.9	98.2	99.1	100.2	97.2	100.7	99.1	100.5
5	129.3	95.4	108.6	99.0	97.6	100.7				
6	126.9	96.0	107.7	99.3	97.0	100.9				
7	125.7	96.3	107.3	99.4	96.8	100.9				
8	125.1	96.4								
9	124.9	96.5								
Dif.	0.27%	-0.07%	0.37%	-0.12%	0.19%	-0.07%	0.40%	-0.30%	0.01%	-0.04%
Dis.	Uniform						Log-Normal			
n_0	10 000	8 500	4 000	15 000	11 000					
n_r	10 000	11 500	16 000	5 000	9 000					
Estimates expressed in % of the true values										
j	s_0	s_r	s_0	s_r	s_0	s_r	s_0	s_r	s_0	s_r
0	166.8	104.3	151.6	105.0	112.8	110.3	108.2	137.5	104.7	163.7
1	112.3	91.3	108.1	92.4	100.1	96.2	101.0	90.7	100.5	90.0
2	100.3	98.6	100.3	98.8	99.2	99.3	99.7	98.0	99.7	97.9
3	98.1	99.8	99.2	99.7	99.1	99.6	99.5	99.2	99.5	99.3
4	97.8	100.0	99.0	99.8	99.1	99.6	99.5	99.4	99.5	99.6
Dif.	0.37%	-0.20%	0.17%	-0.12%	0.00%	-0.01%	0.04%	-0.19%	0.03%	-0.25%

Table S2 shows that in the iterative calculations with extremely high overvaluation of the zeroth s_0 estimates, see the first three pairs of estimates obtained from the normally distributed datasets, the number of iterations needed to obtain a stable final estimate was higher than in the calculations with lower overvaluations, when only 4 or fewer iterations were sufficient. For example, in the computation with the highest overvaluation, 10 000/10 000 variant, 9 iterations had to be used. It should be noted that the number of iterations needed depended

on which of the two zeroth parameter estimates was used at the beginning of the successive estimation steps. When the zeroth s_r estimate with a minor overvaluation was applied in the first corrective step, the stable final s_0 estimate, 124.3%, was obtained after 7 iterative steps. However, when processing empirical datasets, it is not known at the beginning of the estimating process, which of these two parameters will have a more overvalued zeroth estimate.

Text S3.2 The variability, bias, and errors of the estimates

The bias of the estimates was evaluated according to the RME values of the means of the pairs of parallel final estimates summarized in Table S3 (ESM_1.pdf), the individual final estimates see Table S4 (ESM_1.pdf). The RME values were compared with the half-widths of the confidence intervals of the s_0 and s_r parameters. In all cases, the half-widths were higher than the RME values, in most cases seven or more times. With such relatively wide confidence intervals, the biases found cannot be considered as statistically significant, even though the investigated means might not have normal distribution in some cases. However, in two cases, the values of the half-width and RME were almost the same. They were the s_r estimates obtained from the datasets with uniformly and log-normally distributed concentrations by using the variants labelled 8 500/11 500 and 15 000/5 000 with RME -0.24% and 0.69%, respectively. Both these RME values, perhaps even *statistically* significant, can be regarded as acceptably low and *practically* insignificant. The errors of the s_0 and s_r estimates were negatively correlated with each other, as proved by a significant value of the Spearman correlation coefficient, -0.853. Since the errors were positive and negative, this meant that a higher positive estimation error of one parameter brought about a greater negative estimation error of the other parameter and vice versa. In this case, there was a problem with the correct evaluation of the significance level of this test because there were only 8 simulated datasets that were processed by a few variants of the proposed procedure, and the coefficient was calculated from 20 couples of the errors.

The variability of the estimates was evaluated according to the differences between the pairs of the parallel *final* estimates (see Table S3). The largest differences were found in the case of the s_0 estimates obtained by processing the normally distributed datasets, over 10% or even tens of percent, and then by processing the datasets distributed uniformly, differences of several percent. The differences of the s_r estimates were substantially lower. They amounted to a few percent and tenths of a percent or less in the case of the normally and uniformly distributed datasets, respectively. In the case of the log-normally and exponentially distributed datasets, the differences between the pairs of the s_0 and s_r estimates were essentially comparable. The highest one reached to 2.4%.

Table S3 also presents the means and differences of the pairs of parallel *zeroth* estimates. It can be noted that the final estimates with large differences between the parallel results, i.e., with high variability, appeared when the zeroth estimates were severely overvalued. This is particularly evident in the cases of the s_0 estimates obtained from the normally distributed datasets by the variants labelled 10 000/10 000, 4 000/16 000, and 2 600/17 400, where the zeroth estimates amounted to 200% of the true value or more, and the differences between the pairs of the parallel final estimates were extremely high. The pair of the s_0 estimates obtained using the 10 000/10 000 variant can be presented as the most striking example. The mean of the zeroth estimates represents 250.1% of the true value, the difference between them is only 5.1%. The mean of the final estimates represents 98.2%, the difference is much larger, 53.4%.

It can be seen that the differences between the pairs of the *zeroth* estimates were in most cases less than the differences between the pairs of the *final* estimates: 8 and all 10 differences for the presented s_r and s_0 estimates, respectively. The pair of the estimates obtained from the normally distributed datasets by the variant 400/19 600

is distinguished by a high difference between the parallel zeroth estimates of s_0 ; its value is 9.8%, while the other differences were less than 6%.

Table S3. Comparison of the pairs of zeroth and final estimates of s_0 and s_r obtained by different variants of the proposed procedure from the couples of the datasets with 20,000 duplicated results simulated with four chosen types of the concentration distribution; the values are expressed in percentages of the true parameter values; Mean and Difference (absolute values) of the 1st and 2nd estimate; RME – relative mean error; HW – half-width of the confidence interval for the mean; n_0 and n_r – numbers of duplicates for calculation of s_0 and s_r ; c – concentration at the boundary between the subsets of the n_0 and n_r duplicated results, ACP – average correction proportion in percentages of the sum of d_i^2 or d_{ri}^2

Distribution		Normal						Exponential			
n_0		10 000		4 000		2 600		400		10 000	
n_r		10 000		16 000		17 400		19 600		10 000	
c		6		4.74		4.31		2.92		1.63	
Iteration number		9		7		7		4		4	
Estimation		s_0	s_r	s_0	s_r	s_0	s_r	s_0	s_r	s_0	s_r
Zeroth	Mean	250.1%	105.1%	209.2%	106.5%	192.8%	107.0%	145.6%	107.8%	107.3%	124.6%
	Difference	5.1%	2.8%	5.6%	1.0%	4.7%	0.66%	9.8%	0.40%	0.09%	0.49%
Final	Mean	98.2%	100.4%	94.5%	101.3%	87.9%	102.3%	89.5%	102.0%	99.5%	99.6%
	Difference	53.4%	7.7%	25.7%	3.9%	17.9%	2.6%	15.5%	2.5%	0.78%	1.9%
	RME	-1.8%	0.4%	-5.5%	1.3%	-12.1%	2.3%	-10.5%	2.0%	-0.46%	-0.42%
	HW	339.2%	49.1%	163.2%	24.6%	113.8%	16.7%	98.4%	15.6%	5.0%	11.8%
ACP		83.6%	8.8%	79.4%	9.5%	79.1%	8.7%	62.2%	10.5%	13.9%	36.1%
Distribution		Uniform						Log-normal			
n_0		10 000		8 500		6 000		11 000		15 000	
n_r		10 000		11 500		14 000		9 000		5 000	
c		5		4.25		3		1.09		1.60	
Iteration number		4		4		4		4		4	
Estimation		s_0	s_r	s_0	s_r	s_0	s_r	s_0	s_r	s_0	s_r
Zeroth	Mean	167.1%	104.3%	151.6%	105.1%	127.6%	107.3%	105.0%	163.6%	108.4%	137.6%
	Difference	0.58%	0.03%	0.000%	0.21%	1.5%	0.23%	0.69%	0.12%	0.45%	0.23%
Final	Mean	99.7%	99.8%	100.1%	99.8%	99.1%	100.0%	100.0%	98.3%	99.8%	99.3%
	Difference	3.9%	0.42%	2.1%	-0.05%	3.7%	0.45%	1.1%	2.4%	0.63%	0.15%
	RME	-0.27%	-0.19%	0.060%	-0.24%	-0.88%	0.009%	0.017%	-1.7%	-0.21%	-0.69%
	HW	24.7%	2.7%	13.4%	0.31%	23.6%	2.8%	6.8%	15.4%	4.0%	0.93%
ACP		64.3%	8.3%	56.5%	9.6%	39.1%	13.1%	9.4%	63.3%	15.1%	47.6%

The greatest difficulties in estimating occurred in the processing of the normally distributed datasets, as the proportion of the d_i difference suitable for the s_0 estimation was extremely low. Only 0.5% of all differences in the datasets belonged to the concentration range, where $s_0^2 > s_r^2 c^2$. When processing data by variants 10 000/10 000, 4 000/16 000, and 2 600/17 400, the s_0 estimates were calculated from differences in the concentration ranges, where the variance $s_r^2 c^2$ was higher than the variance s_0^2 even up to about 8, 5 and 4 times; the average correction proportions accounted for 83.6%, 79.4% and 79.1% of $\sum d_i^2$, respectively. This resulted in very overvalued zeroth estimates of s_0 and high variability of the final estimates of s_0 and consequently also s_r estimates (see Table S3). In the above-mentioned case of variant 400/19 600, the n_0 value was reduced in such a way that the average correction

proportion accounted for about 62,2%. With such an excessively low value of n_0 , the variability of values of $\sum d_i^2$ was high compared to the cases where n_0 represented several thousand, and therefore the variability of the zeroth estimates was so high, see above.

Table S4. Pairs of the zeroth and final estimates of s_0 and s_r obtained by chosen variants of the proposed procedure from the couples of the dataset with 20 000 duplicate results simulated with chosen types of the concentration distribution; the values are expressed in percentages of the true parameter values; n_0 and n_r – numbers of duplicates used for calculation of s_0 and s_r , Correction proportion – percentages of the correction from the sum of d_i^2 or d_{ri}^2

Distribution		Normal						Exponential			
n_0/n_r		10 000/10 000		40 000/16 000		26 000/17 400		400/10 600		10 000/10 000	
Estimate of		s_0	s_r	s_0	s_r	s_0	s_r	s_0	s_r	s_0	s_r
Zeroth estimate	1 st	252.6%	103.7%	212.0%	106.0%	195.2%	106.7%	150.5%	107.6%	107.2%	124.9%
	2 nd	247.5%	106.5%	206.4%	107.0%	190.5%	107.3%	140.7%	108.0%	107.3%	124.4%
Final estimate	1 st	124.9%	96.5%	107.3%	99.4%	96.8%	100.9%	97.2%	100.7%	99.1%	100.5%
	2 nd	71.5%	104.2%	81.6%	103.3%	78.9%	103.6%	81.8%	103.2%	99.9%	98.7%
Correction proportion	1 st	75.6%	13.4%	74.4%	12.1%	75.4%	10.4%	58.3%	12.3%	14.5%	35.2%
	2 nd	91.7%	4.2%	84.4%	6.9%	82.8%	6.9%	66.2%	8.7%	13.3%	37.1%
Distribution		Uniform				Log-normal					
n_0/n_r		10 000/10 000		8 500/11 500		6 000/14 000		11 000/9 000		15 000/5 000	
Estimate of		s_0	s_r	s_0	s_r	s_0	s_r	s_0	s_r	s_0	s_r
Zeroth estimate	1 st	166.8%	104.3%	151.6%	105.0%	126.8%	107.2%	104.7%	163.7%	108.2%	137.5%
	2 nd	167.4%	104.3%	151.6%	105.2%	128.3%	107.4%	105.4%	163.6%	108.6%	137.7%
Final estimate	1 st	97.8%	100.0%	99.0%	99.8%	97.3%	100.2%	99.5%	99.6%	99.5%	99.4%
	2 nd	101.7%	99.6%	101.1%	99.7%	101.0%	99.8%	100.6%	97.1%	100.1%	99.2%
Correction proportion	1 st	64.9%	8.1%	56.9%	9.4%	39.5%	12.7%	9.2%	62.8%	15.0%	47.4%
	2 nd	63.7%	8.5%	56.0%	9.8%	38.7%	13.5%	9.5%	63.7%	15.2%	47.7%

It is clear that the proposed procedure is unsuitable for processing the datasets simulated with the chosen s_0 and s_r values and the given normal distribution of concentrations. However, in a case of processing similar datasets, but with a realistic number of d_i differences, i.e., maximally hundreds, this procedure would probably not be used at all, because in the plot $|d_i|$ vs. $\bar{c}_i$ there would be practically no points in the concentration range with a constant variance of the differences, the probability of their occurrence is 1 in 200. Then the model with a constant RSD would probably be considered, i.e., estimating RSD from d_{ri} values using Eq. (3).

In the case of the log-normally distributed datasets, there were similar difficulties as in the processing of the normally distributed datasets, however, with the s_r estimates and less serious. The zeroth estimates of s_r had the highest overvaluation. Only about 14% of all d_i differences of the datasets belonged to the concentration range, where $s_r^2 c^2 > s_0^2$. When estimating s_r using variants 11 000/9 000 and 15 000/5 000 some of the d_{ri} differences included in the calculation were from the concentration ranges, where the s_0^2 variance was higher than the $s_r^2 c^2$ variance up to about 4 and 2 times. the average correction proportions accounted for 63.3% and 47.6%, respectively. The relative difference between the parallel final s_r estimates was -2.4% in the first case, and only -0.15% in the second case.

When processing the uniformly distributed datasets, only ca. 21% of all d_i differences of the datasets belonged to the concentration range where $s_0^2 > s_r^2 c^2$. By the variants 10 000/10 000, 8 500/11 500, and 6 000/14 000, the s_0 estimates were calculated from the differences in the concentration ranges, where $s_r^2 c^2$ was almost up to 5.5 times, 4 times and 2 times higher than s_0^2 . The average correction proportions accounted for 64.3%, 56.5% and 39.1%, respectively. In the first and second cases, the differences between the two parallel final s_0 estimates were 3.9% and 2.1%. In the third case, where practically only the d_i differences suitable for estimating s_0 were included, the difference between the final estimates was surprisingly relatively large, 3.7%; however, a relatively high difference was already between the zeroth estimates, 1.5% versus 0.58% and almost 0% (see Table S3).

When processing the exponentially distributed datasets, approximately 60% of all d_i differences of the datasets belonged to the concentration range, where $s_0^2 > s_r^2 c^2$. When $n_0 = n_r = 10\,000$ was chosen, the average correction proportions in estimating s_0 and s_r were low: 13.9% and 36.1%, respectively. It can be seen that slightly worse conditions were when estimating s_r . The relative differences between the parallel final estimates of s_0 and also s_r were low: 0.8% and 1.9% respectively.

The results show that the inclusion of an extremely high proportion of d_i or d_{ri} differences unsuitable for estimating a given parameter, as well as a substantial reduction in the number of differences processed, cause a higher variability level of the final estimates of parameters.

Text S4 Overview of the estimates obtained from datasets with $n = 200$ and some their statistics

This section contains tables with individual estimates of s_0 and s_r obtained from the tens of datasets with $n = 200$ simulated with selected concentration distributions – see Tables S5 and S6 (ESM_1.pdf). The datasets were processed by several variants of the proposed procedure and also by various regression procedures.

Table S5. Estimates of s_0 and s_r obtained by chosen variants of the proposed procedure and regression procedures from 10 datasets with 200 duplicates simulated with the **uniform** concentration distribution; estimates expressed in % of the true values

Procedure	Parameter	Order number of the simulated dataset									
		1	2	3	4	5	6	7	8	9	10
Proposed $c = 2.6$	s_0	79.8%	99.9%	90.7%	62.7% ^m	106.3%	70.4%	129.7% ^M	97.2%	107.0%	68.7%
	s_r	95.1% ^m	107.7%	106.1%	108.3% ^M	97.8%	97.8%	97.0%	96.2%	100.4%	97.0%
Proposed $c = 4$	s_0	96.9%	91.9%	81.4%	104.5%	102.4%	80.3% ^m	93.9%	115.5% ^M	81.5%	93.4%
	s_r	90.8%	109.3% ^M	108.0%	97.9%	97.8%	96.6%	108.7%	90.3%	106.3%	87.0% ^m
Proposed 85/150	s_0	90.6%	96.9%	65.8% ^m	106.9%	93.9%	80.2%	105.2%	113.2% ^M	67.6%	101.1%
	s_r	94.0%	108.0%	110.3% ^M	97.4%	103.0%	95.4%	101.9%	91.2%	103.5%	87.9% ^m
Proposed 100/100	s_0	105.8%	129.3% ^M	21.8% ^m	121.8%	96.1%	96.1%	109.7%	107.6%	70.8%	94.3%
	s_r	88.7%	100.2%	114.3% ^M	93.1%	99.2%	92.6%	105.4%	92.7%	108.3%	86.7% ^m
RMS WLS	s_0	73.7% ^m	107.4%	104.9%	80.7%	99.1%	87.1%	123.0% ^M	102.6%	112.6%	82.0%
	s_r	95.8%	107.0% ^M	102.0%	104.5%	99.9%	93.7%	101.9%	94.9%	100.6%	90.3% ^m
MAD WLS	s_0	51.8% ^m	84.8%	90.4%	94.2%	104.2%	97.9%	138.6% ^M	114.9%	135.0%	94.6%
	s_r	112.7%	128.1% ^M	110.3%	92.9%	98.1%	90.2%	94.2%	113.2%	103.0%	84.1% ^m
RMS OLS	s_0	159.0% ^M	73.5%	93.3%	93.6%	108.3%	103.0%	146.3%	107.8%	0.0% ^{*m}	118.8%
	s_r	72.7% ^m	112.6%	105.2%	100.4%	97.5%	91.9%	97.1%	93.5%	114.2% ^M	78.3%
MAD OLS	s_0	188.7%	112.6%	0.0% ^{*m}	0.0% ^{*m}	74.4%	138.8%	190.3% ^M	159.0%	0.0% ^{*m}	95.4%
	s_r	75.3% ^m	119.9%	118.5%	107.1%	102.4%	80.9%	76.1%	101.1%	128.5% ^M	82.7%

* 0 was given as the s_0 estimate when the s_0^2 estimate obtained by regression procedure was negative and insignificant; ^m minimum; ^M maximum

Table S6. Estimates of s_0 and s_r obtained by chosen variants of the proposed procedure and regression procedures from 10 dataset with 200 duplicates simulated with the *exponential* concentration distribution; estimates expressed in % of the true values

Procedure	Parameter	Order number of the simulated dataset									
		1	2	3	4	5	6	7	8	9	10
Proposed 110/90	s_0	92.4%	82.1% ^m	104.0%	96.6%	99.3%	92.0%	104.5%	106.8%	97.1%	111.9% ^M
	s_r	109.1%	105.3%	90.7%	105.2%	106.2%	120.7% ^M	116.8%	102.4%	90.8%	90.3% ^m
Proposed 100/100	s_0	95.5%	84.7% ^m	106.7%	95.0%	96.9%	94.6%	111.4% ^M	105.5%	100.3%	109.1%
	s_r	104.6%	101.6%	83.7% ^m	107.0%	110.1%	117.3% ^M	108.5%	104.0%	83.9%	94.6%
Proposed 100/120	s_0	97.9%	85.4% ^m	105.2%	94.0%	95.5%	97.5%	112.9% ^M	104.2%	100.9%	110.7%
	s_r	89.4%	98.2%	96.5%	112.2%	118.2% ^M	100.2%	101.6%	111.8%	79.7% ^m	86.7%
Proposed 100/140	s_0	98.0%	84.6% ^m	104.4%	94.8%	100.9%	95.0%	110.4%	104.7%	102.7%	111.6% ^M
	s_r	89.6%	103.9%	103.5%	109.9%	87.6%	117.0% ^M	115.1%	110.3%	63.8% ^m	82.4%
RMS WLS	s_0	103.4%	81.8% ^m	99.8%	94.5%	100.8%	100.1%	118.7% ^M	106.9%	97.8%	110.4%
	s_r	96.5%	108.8%	101.2%	106.1%	97.3%	110.2% ^M	93.3% ^m	100.0%	95.8%	94.5%
MAD WLS	s_0	105.4%	82.0% ^m	92.4%	85.0%	127.6% ^M	98.5%	115.1%	113.6%	126.3%	103.8%
	s_r	92.8%	118.6% ^M	100.3%	118.2%	89.4% ^m	107.3%	115.7%	99.8%	95.4%	109.5%
RMS OLS	s_0	121.9%	33.1% ^m	101.1%	93.9%	109.8%	99.2%	129.4% ^M	108.3%	76.4%	114.6%
	s_r	80.2% ^m	124.9% ^M	101.4%	105.7%	89.5%	110.6%	82.3%	96.5%	111.9%	92.6%
MAD OLS	s_0	115.2%	61.4%	47.5% ^m	84.3%	132.5% ^M	115.8%	118.9%	95.2%	127.4%	105.8%
	s_r	87.2%	127.3% ^M	125.7%	115.8%	85.7% ^m	93.1%	110.7%	111.7%	93.8%	110.4%

^m minimum; ^M maximum

The statistical characteristics (mean, RSD and others) of the sets of s_0 and s_r estimates presented in Tables S5 and S6 (ESM_1.pdf) are summarized in the main paper, see Tables 1 and 2. The distributions of these sets of estimates were studied on the box plots in Fig. S2 (ESM_1.pdf). These plots depict the medians, lower and upper quartiles, and non-outlier ranges defined by the inner fences. They also identify mild and extreme outliers and allow visual comparisons of the variability and skewness of these sets. These sets of s_0 and s_r estimates were also tested for homogeneity and normality (see Subchapter “Checking homogeneity and normality of the estimate sets” and Text S1.2 in ESM_1.pdf). The statistics found are summarised in Tables S7 and S8 (ESM_1.pdf). Since the processed datasets were only simulated, the estimates identified as outliers could not appear due to the gross experimental errors. They could represent some anomalies of the estimate set distributions, which were caused by the investigated estimation procedures; however, they could also occur as a consequence of the type I error.

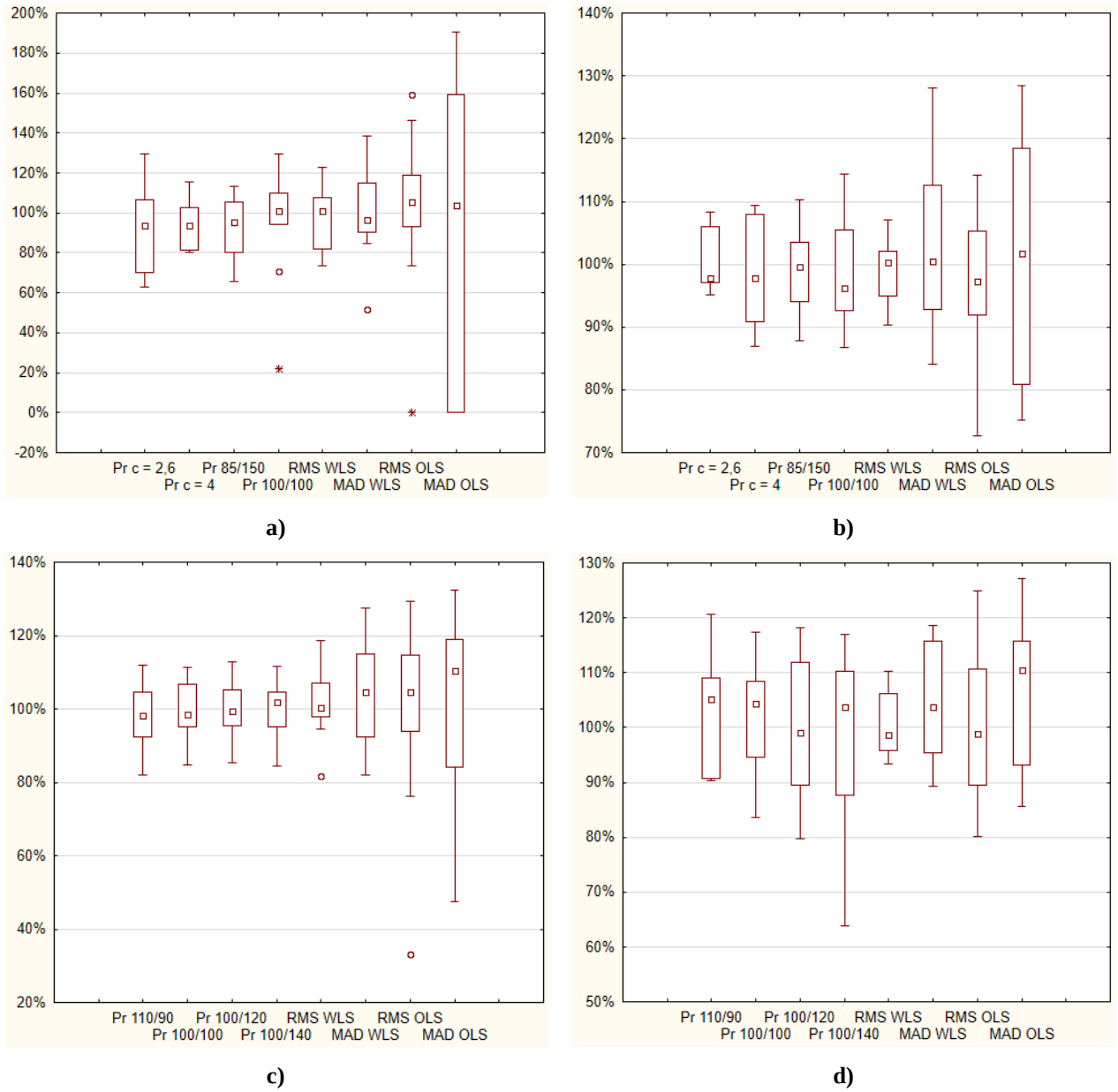


Fig. S2. Box plots of the sets of the estimates expressed in percentages of the true values, the sets were obtained by marked procedures from the datasets with $n = 200$; the estimates of (a) s_0 , (b) s_r from the **uniformly** distributed datasets; and (c) s_0 , (d) s_r from the **exponentially** distributed datasets; □ median; □ range between the upper and lower quartile; ┘ non-outlier range; ○ mild outlier; * extreme outlier

The box plots in Fig. S2 clearly show, e.g., that the dispersion of s_r estimates from both types of distributions obtained by the RMS WLS procedure significantly lower than the s_r estimates obtained by other procedures, except for estimates obtained by the $c = 2.6$ variant from uniformly distributed datasets. It is also clearly visible that estimates obtained by the proposed procedure, when a large proportion of unsuitable differences have been included in the estimation, suffer from large negative errors: (i) in the s_0 estimates obtained from uniformly distributed datasets by the 100/100 variant, there are two outliers, one of them extreme; (ii) s_0 estimates obtained from exponentially distributed datasets by the 100/140 variant are strongly skewed to low values. The diagrams show that the s_r estimates obtained from exponentially distributed datasets by the proposed procedure have a relatively large variance; that estimates of both parameters using OLS regression procedures have a relatively high variance or contain outliers.

Table S7. Statistics of the s_0 and s_r estimate sets obtained by the chosen variants of the proposed procedure and by the regression procedures from 10 datasets with 200 duplicates simulated with the **uniform** concentration distribution – see Table S5; G_m , G_M and G_2 – statistics of the Grubbs test for a single outlier at the left and right end and for a pair of outliers at opposite ends of the set, respectively; D – statistics of the Kolmogorov-Smirnov/Lilliefors test for normality

Procedure	Parameter	Statistics					
		Skewness	Kurtosis	G_m	G_M	G_2	D
Proposed $c = 2.6$	s_0	0.30	-0.46	1.36	1.83	3.20	0.140
	s_r	0.82	-1.19	1.04	1.58	2.62	0.289 *
Proposed $c = 4$	s_0	0.44	-0.22	1.22	1.88	3.11	0.168
	s_r	-0.08	-1.63	1.47	1.21	2.67	0.200
Proposed 85/150	s_0	-0.62	-0.77	1.62	1.30	2.92	0.162
	s_r	-0.01	-1.00	1.56	1.52	3.07	0.140
Proposed 100/100	s_0	-1.72 *	3.68 *	2.42 *	1.12	3.53	0.287 *
	s_r	0.55	-0.70	1.27	1.80	3.06	0.210
RMS WLS	s_0	-0.01	-1.02	1.49	1.62	3.11	0.144
	s_r	-0.21	-0.78	1.67	1.53	3.20	0.166
MAD WLS	s_0	-0.21	0.77	1.95	1.51	3.46	0.164
	s_r	0.52	-0.22	1.39	1.90	3.29	0.137
RMS OLS	s_0	-1.24	2.98 *	2.32 *	1.35	3.67	0.235
	s_r	-0.50	-0.17	1.78	1.34	3.12	0.170
MAD OLS	s_0	-0.22	-1.51	1.27	1.25	2.51	0.197
	s_r	0.08	-1.57	1.23	1.49	2.72	0.201

* indication of the skewness and kurtosis values that contradict the assumption of normality and indication of the test statistics higher than the critical values at $\alpha = 0.05$ and $n = 10$

The Grubbs tests identified outliers only in the s_0 estimate sets, all at $\alpha = 0.05$, essentially in accordance with the box plots (see Fig. S2): the minimum in the set obtained by the 100/100 variant from the datasets uniformly distributed, the minima in the sets obtained by the RMS OLS procedure from the datasets with both distribution types; furthermore, in the case of the set obtained by the RMS WLS procedure from the exponentially distributed datasets, a pair of outliers at opposite ends was found (see values G_m and G_2 in Tables S7 and S8). These tests detected outliers under the assuming normally distributed sets; the critical values of these test were 2.29 and 3.68 when testing a single outlier and range, respectively, at $\alpha = 0.05$ and $n = 10$ (see [4]).

Furthermore, the skewness and kurtosis coefficients of the estimate sets were evaluated. Values inconsistent with normality, i.e., with absolute values twice higher than the standard errors of the coefficient estimates, were found only for three s_0 estimate sets with the outlying minima: the sets obtained by the RMS OLS procedure from the datasets with both distribution types and the set obtained by the 100/100 variant from the uniformly distributed data sets. All these three sets have too high values of the kurtosis coefficient. The sets obtained by the 100/100 variant and the RMS OLS procedure from the uniformly and exponentially distributed datasets, respectively, in addition have too negative values of the skewness coefficient. In the case of normally distributed sets with $n = 10$, the standard errors equal 0.687 and 1.33 for skewness and kurtosis, respectively [2s - Text s8 in ESM_1.pdf].

The Kolmogorov-Smirnov/Lilliefors test of normality identified only two sets of estimates as samples not coming from normally distributed populations. Both sets were obtained by the proposed procedure; the critical D value at $\alpha = 0.05$ and $n = 10$ equals 0.262 [1s - Text s8 in ESM_1.pdf]. The former set was again the set of the s_0 estimates obtained by the 100/100 variant from the uniformly distributed datasets. The latter set was, somewhat surprisingly, the set of the s_r estimates that was obtained by the variant $c = 2.6$ from the uniformly distributed datasets, even though all problematic cases mentioned above were only the s_0 estimates sets. None of the previous statistical tests used indicated a deviation of the distribution of this set from normality; however, the box plot (see Fig. S2b) shows that the difference between the median and the lower quartile is too small.

Table S8. Statistics of the s_0 and s_r estimate sets obtained by the chosen variants of the proposed procedure and by the regression procedures from 10 datasets with 200 duplicates simulated with the *exponential* concentration distribution – see Table S6; G_m , G_M and G_2 – statistics of the Grubbs test for a single outlier at the left and right end and for a pair of outliers at opposite ends of the set, respectively; D – statistics of the Kolmogorov-Smirnov/Lilliefors test for normality

Procedure	Parameter	Statistics					
		Skewness	Kurtosis	G_m	G_M	G_2	D
Proposed 110/90	s_0	-0.41	0.25	1.92	1.54	3.46	0.130
	s_r	0.05	-0.89	1.27	1.59	2.86	0.189
Proposed 100/100	s_0	-0.31	-0.33	1.86	1.40	3.26	0.155
	s_r	-0.66	-0.33	1.62	1.43	3.05	0.201
Proposed 100/120	s_0	-0.16	0.00	1.83	1.52	3.36	0.122
	s_r	-0.01	-0.76	1.62	1.54	3.16	0.145
Proposed 100/140	s_0	-0.60	0.54	2.00	1.36	3.36	0.139
	s_r	-0.91	0.22	2.03	1.10	3.13	0.220
RMS WLS	s_0	-0.28	1.51	2.00	1.76	3.77*	0.157
	s_r	0.62	-1.12	1.17	1.62	2.79	0.194
MAD WLS	s_0	0.02	-1.09	1.45	1.42	2.87	0.110
	s_r	0.03	-1.53	1.43	1.30	2.73	0.158
RMS OLS	s_0	-1.66 *	3.39 *	2.39 *	1.11	3.51	0.229
	s_r	0.29	-0.50	1.37	1.80	3.17	0.091
MAD OLS	s_0	-0.90	-0.23	1.87	1.13	3.01	0.200
	s_r	-0.05	-1.47	1.34	1.39	2.72	0.210

* indication of the skewness and kurtosis values that contradict the assumption of normality and indication of the test statistics higher than the critical values at $\alpha = 0.05$ and $n = 10$

Text S5 Correction proportion from the sum of the squared differences

In Subchapter “Correction proportion from the sum of squared differences: its purpose and calculation”, a correction proportion from the sum of squared differences has been introduced, which draws attention to the inclusion of a large proportion of differences unsuitable for estimating a given parameter. In this text, the values of this parameter are monitored for estimates (see Table S4 and Tables S5 and S6) obtained from simulated datasets in relation to the errors of individual estimates and the variability of the sets of these estimates.

In the case of the parameter estimates obtained from datasets with $n = 20\,000$, a high correlation was proved for the s_0 estimates: the larger the average correction proportions in the estimation, the greater the absolute values of the differences between the parallel estimates of the first and second series; i.e., the estimates have greater variability (see data in Table S3, ESM_1.pdf); Spearman’s correlation coefficient was equal to 0.927. This relationship is clearly evident in 6 cases with a large proportion of correction, greater than 55%. For the s_r estimates this correlation was not proved, the correlation coefficient was equal to -0.297. However, there were only two cases with a higher correction proportion, this was when estimating from the log-normally distributed datasets.

In the case of those s_r estimates, a different relationship proved to be significant: the values of the s_r estimates were negatively correlated with the individual values of the correction proportions (see data in Table S4, ESM_1.pdf). Spearman’s correlation coefficient was -0.737. This relationship can be assumed on the basis of the equations for parameter estimation (Equations (14) and (15)), because random overestimation of the correction term causes an underestimation of the parameter, and vice versa. This relationship was not proved for s_0 estimates; Spearman’s correlation coefficient was -0.421. The critical values of Spearman’s coefficient at $\alpha = 0.05$ are 0.648 and 0.447 for n equal to 10 and 20, respectively; as stated in Text S3.2, however, these critical values are problematic in the case of these assessments.

The first relationship examined, that a larger correction proportion is correlated with a greater variance of the estimates, requires a precise estimation of the correction proportions and can therefore be better demonstrated by the average values. The second relationship examined, which mainly monitors the random errors, is better manifested in the evaluation of individual values of the correction proportion.

Table S9. Average correction proportions (ACP) and RSD values for the sets of s_0 and s_r estimates obtained by the chosen variants of the proposed procedure from the 10 datasets with 200 duplicates simulated with the uniform and exponential concentration distribution; the pairs of values confirming an increase in RSD with ACP are highlighted

Dataset distribution	Variant	s_0		s_r	
		ACP	RSD	ACP	RSD
Uniform	$c = 2.6$	37.3%	21.0%	13.2%	5.1%
	$c = 4$	56.5%	11.4%	9.8%	8.4%
	85/150	58.8%	16.3%	15.2%	7.3%
	100/100	63.2%	30.4%	9.3%	9.0%
Exponential	110/90	19.6%	8.6%	30.3%	10.6%
	100/100	15.1%	8.2%	35.9%	11.0%
	100/120	14.4%	8.2%	47.0%	12.2%
	100/140	13.8%	8.0%	58.3%	17.0%

These relationships were also investigated in the case of those estimates obtained from the datasets with 200 duplicates. Table S9 (ESM 1.pdf) summarises the average correction proportions and RSD values for the s_0

and s_r estimate sets calculated by different variants of the procedure from 10 datasets simulated with a certain concentration distribution. The results show that if the average correction proportions was higher than 30%, for s_0 and s_r estimates obtained from datasets distributed uniformly and exponentially, respectively, the RSD values increased with the average correction proportion. Only the set of estimates s_0 obtained by the variant $c = 2.6$ was exceptional, because for the average correction proportion of 37.3%, the RSD = 21% found was too high compared to the other cases. It should be noted that in Subchapter “The estimates obtained by the proposed procedure” it was stated for this set that this high variability of the final estimates was mainly due to the variability of the $\sum d_i^2$ values, since the zeroth estimates already had a surprisingly high variability.

Table S10. Correction proportions, CP, and relative errors, RE, of s_0 and s_r estimates obtained by variants of the proposed procedure from 10 datasets with 200 duplicate results simulated with the **uniform** concentration distribution; SpCC – Spearman’s correlation coefficient; pairs of values mentioned in the text are highlighted

Parameter			S_0					
Variant	$c = 2.6$		$c = 4$		85/150		100/100	
Dataset	CP	RE	CP	RE	CP	RE	CP	RE
1 st	41.3%	-20.2%	53.1%	-3.1%	59.0%	-9.4%	56.5%	5.8%
2 nd	33.5%	-0.1%	62.8%	-8.1%	60.2%	-3.1%	50.2%	29.3%
3 rd	37.6%	-9.3%	69.4%	-18.6%	79.8%	-34.2%	98.1%	-78.2%
4 th	55.5%	-37.3%	46.7%	4.5%	44.6%	6.9%	45.8%	21.8%
5 th	27.3%	6.3%	51.2%	2.4%	59.9%	-6.1%	64.8%	-3.9%
6 th	48.0%	-29.6%	64.7%	-19.7%	64.6%	-19.8%	61.8%	-3.9%
7 th	23.1%	29.7%	63.1%	-6.1%	57.0%	5.2%	63.7%	9.7%
8 th	30.2%	-2.8%	40.9%	15.5%	43.6%	13.2%	55.3%	7.6%
9 th	29.3%	7.0%	65.2%	-18.5%	79.0%	-32.4%	83.8%	-29.2%
10 th	46.7%	-31.3%	47.7%	-6.6%	40.0%	1.1%	52.1%	-5.7%
SpCC	-0.964 *		-0.867 *		-0.879 *		-0.709 *	

Parameter			S_r					
Variant	$c = 2.6$		$c = 4$		85/150		100/100	
Dataset	CP	RE	CP	RE	CP	RE	CP	RE
1 st	11.8%	-4.9%	11.5%	-9.2%	14.7%	-6.0%	11.7%	-11.3%
2 nd	13.4%	7.7%	7.7%	9.3%	13.6%	8.0%	13.8%	0.2%
3 rd	12.0%	6.1%	6.2%	8.0%	6.1%	10.3%	0.4%	14.3%
4 th	5.7%	8.3%	11.5%	-2.1%	21.3%	-2.6%	14.0%	-6.9%
5 th	16.5%	-2.2%	11.2%	-2.2%	13.7%	3.0%	8.0%	-0.8%
6 th	8.4%	-2.2%	7.5%	-3.4%	11.9%	-4.6%	9.3%	-7.4%
7 th	23.6%	-3.0%	8.0%	8.7%	16.0%	1.9%	8.9%	5.4%
8 th	15.1%	-3.8%	15.7%	-9.7%	23.1%	-8.8%	11.0%	-7.3%
9 th	16.5%	0.4%	6.4%	6.3%	6.7%	3.5%	3.7%	8.3%
10 th	8.9%	-3.0%	12.5%	-13.0%	25.2%	-12.1%	11.7%	-13.3%
SpCC	-0.224		-0.673 *		-0.794 *		-0.600	

* indication of the coefficients higher than 0.648, which is the critical value of Spearman’s coefficient at $\alpha = 0.05$ and $n = 10$ (see [26])

When evaluating the individual correction proportion from those datasets with only 200 duplicates, it was again proved that subtraction of a randomly overestimated term causes an underestimation of the final estimate of the parameter and vice versa. Tables S10 and S11 (ESM_1.pdf) present the individual correction proportions, relative errors, RE, of the estimates expressed in percentages of the true values, and Spearman’s coefficients found.

Significant values of these coefficients were again found especially when the s_0 and s_r estimates were obtained from those uniformly and exponentially, respectively, distributed datasets. In these cases, larger proportions of less suitable and unsuitable differences were included in the estimations and the average correction proportions were above 30% (see Table S9). Table S11 shows that an insignificant correlation coefficient was found only for the s_r estimates obtained by the variant 100/100 from the exponentially distributed datasets. The correlation was also significant for the two s_r estimates from uniformly distributed datasets, when the average correction proportions were low.

Table S11. Correction proportions, CP, and relative errors, RE, of s_0 and s_r estimates obtained by variants of the proposed procedure from 10 datasets with 200 duplicate results simulated with the *exponential* concentration distribution; SpCC – Spearman’s correlation coefficient; pairs of values mentioned in the text are highlighted

Parameter			S ₀					
Variant	110/90		100/100		100/120		100/140	
Dataset	CP	RE	CP	RE	CP	RE	CP	RE
1 st	21.9%	-7.6%	15.4%	-4.5%	11.2%	-2.1%	11.0%	-2.0%
2 nd	23.1%	-17.9%	16.7%	-15.3%	15.4%	-14.6%	17.0%	-15.4%
3 rd	11.6%	4.0%	7.9%	6.7%	10.4%	5.2%	11.8%	4.4%
4 th	22.4%	-3.4%	18.9%	-5.0%	20.6%	-6.0%	19.4%	-5.2%
5 th	19.0%	-0.7%	17.6%	-3.1%	20.1%	-4.5%	10.8%	0.9%
6 th	25.2%	-8.0%	18.6%	-5.4%	13.4%	-2.5%	17.9%	-5.0%
7 th	25.8%	4.5%	17.0%	11.4%	14.8%	12.9%	18.5%	10.4%
8 th	16.7%	6.8%	14.2%	5.5%	16.3%	4.2%	15.5%	4.7%
9 th	14.6%	-2.9%	9.8%	0.3%	8.8%	0.9%	5.5%	2.7%
10 th	15.8%	11.9%	15.0%	9.1%	12.5%	10.7%	11.1%	11.6%
SpCC	-0.455		-0.503		-0.479		-0.236	
Parameter			S _r					
Variant	110/90		100/100		100/120		100/140	
Dataset	CP	RE	CP	RE	CP	RE	CP	RE
1 st	23.7%	9.1%	30.9%	4.6%	51.6%	-10.6%	63.1%	-10.4%
2 nd	20.5%	5.3%	27.5%	1.6%	41.3%	-1.8%	49.4%	3.9%
3 rd	43.1%	-9.3%	54.0%	-16.3%	56.3%	-3.5%	61.4%	3.5%
4 th	26.2%	5.2%	28.3%	7.0%	35.2%	12.2%	46.3%	9.9%
5 th	32.5%	6.2%	34.1%	10.1%	38.9%	18.2%	65.0%	-12.4%
6 th	22.3%	20.7%	27.8%	17.3%	47.4%	0.2%	50.7%	17.0%
7 th	22.9%	16.8%	31.9%	8.5%	44.7%	1.6%	48.9%	15.1%
8 th	32.4%	2.4%	35.7%	4.0%	41.5%	11.8%	53.5%	10.3%
9 th	39.9%	-9.2%	49.4%	-16.1%	60.6%	-20.3%	80.1%	-36.2%
10 th	39.1%	-9.7%	39.7%	-5.4%	52.6%	-13.3%	64.6%	-17.6%
SpCC	-0.782 *		-0.636		-0.879 *		-0.830 *	

* indication of the coefficients higher than 0.648, which is the critical value of Spearman’s coefficient at $\alpha = 0.05$ and $n = 10$ (see [26])

Tables S10 and S11 show that the variability of the individual correction proportions is large. Thus, the question is whether, with such variability, the correction proportion can be an indicator of the possibility of a large random error in the parameter estimation. Did the correction proportion values correctly warn against the possibility of large random errors, e.g., when estimating s_0 from the uniformly distributed datasets by the 100/100 variant, i.e., when including the largest proportions of d_i differences unsuitable for this estimation? The average

correction proportion was 63.2%, the individual values range from 45.8% to 98.1% (see Tables S9 and S10). The two highest correction proportions found, 98.1% and 83.8%, clearly indicated the possibility of major error. Accordingly, in these cases, the parameter s_0 was estimated with the largest negative errors: -78.2%, extremely outlier minimum estimate, and -29.2%. The two largest positive errors, 29.3% and 21.8%, were associated with two lowest correction proportions, 50.2% and 45.8%, respectively. These values of the proportions, as values individually found, are not alarming, only a comparison with the above average correction proportion shows that they are underestimated. In the case of the 85/150 variant, the two highest correction proportions, 79.8% and 79.0%, again clearly pointed to the estimates with the largest negative errors, -34.2% and -32.4% respectively. The third highest correction proportion, 64.6%, pointed to the third largest negative error, -19.8%. The correction proportion could be considered a warning; however, the error is not extremely alarming. The other values of the correction proportion for the s_0 estimates obtained by the variants 100/100, 85/100 and $c = 4$ were less than 70%. The errors of the obtained s_0 estimates ranged from -20% to 16%; such errors would be acceptable even in the case of s_0 estimates obtained by the RSD WLS procedure.

For the s_0 estimates obtained by the variant $c = 2.6$, the largest errors of the estimates were -37.3%, -31.3%, -29.6%, and 29.7%, respectively, and the corresponding values of the correction proportions were 55.5%, 46.7%, 48.0% and 23.1, respectively. It should be noted that the relatively high variability of the set of these estimates was also due to the higher variability of the $\sum d_i^2$ values (see above). Large negative estimation errors were not directly indicated by any particularly high individual values of the correction proportion. Only a comparison of the individual values of the proportion with the average correction proportion, 37.3%, shows that the large negative and positive errors of the estimates, respectively, were caused by the largest positive and negative deviations of the individual values of the correction proportion from the average.

A similar situation occurred in the estimates of s_r from the exponentially distributed datasets (see Table S11). When estimating by the variant 100/140, which was the variant with the largest included proportion of unsuitable differences, the visibly highest correction proportion of 80.1% pointed to the largest negative estimate error, - 36.2%. The second largest negative error in the table, -20.3%, and the largest for the 100/120 variant, was associated with a correction proportion of 60.6%, the largest for that variant. These cases could be seen as a warning of the possibility of a major negative error. The other high correction proportions of 56.3% and 52.6% did not differ much from the average correction proportion of 47.0%. They were associated with errors of -3.5% and -13.3%. The largest positive error, 20.7%, was associated with a correction proportion of 22.3%, the second smallest proportion for the 110/90 variant. This proportion was undervalued compared to the average of 30.3% (see Table 9).

The analysis of the results summarised in Tables S9, S10 and S11 confirms the conclusions of the analysis of the estimates obtained from the datasets with $n = 20\ 000$. The variability in the estimates s_0 and s_r increases as the correction proportion increases, i.e., as the proportion of unsuitable differences included in the estimate increases. At high values of the correction proportion, its values negatively correlate with errors in the parameter estimation, i.e., a large overestimation of the correction term causes a large negative error in the parameter estimation and vice versa.

However, when estimating a parameter from a given empirical dataset, only one estimate of the pair s_0 and s_r is obtained, i.e., only one random value of the correction proportion for each parameter. If the value of any of the proportion is too high, a large random error of the corresponding parameter estimate can reasonably be

expected. The subset of duplicate differences included in this estimation should therefore be modified. This means that the proportion of differences not suitable for this estimation should be reduced and then the estimation of both parameters should be repeated. If the value of any of the correction proportion is acceptably low, there are two possible interpretations. This may indicate a suitably chosen subset of differences used for the estimation, and then a low random error of the estimate can be assumed. On the other hand, given the possibility of high variability of the correction proportion, it cannot be ruled out that the value of the proportion is severely underestimated, the obtained estimate of the parameter would then be significantly overestimated. It is therefore obvious that the correction proportion can only warn of the possibility of an extreme negative error in the parameter estimate. Unfortunately, it does not warn of the possibility of a large positive error.

On a purely logical basis, 50% can be recommended as the maximum acceptable value for the correction proportion. However, this recommended value should first be verified when estimating s_0 and s_r from empirical datasets with different numbers of duplicates, different values of the estimated parameters and different concentration distributions of the measured samples.

Text S6 Statistical significance of the s_0^2 and s_r^2 estimates obtained by regression procedures

Table S12. P-values of the t -test assessing the significance of the s_0^2 and s_r^2 estimates obtained by the regression procedures from the datasets with two types of the concentration distribution with $n = 200$

Results from the datasets with <i>uniformly</i> distributed concentration								
no.	RMS WLS		MAD WLS		RMS OLS		MAD OLS	
	s_0^2	s_r^2	s_0^2	s_r^2	s_0^2	s_r^2	s_0^2	s_r^2
1	0.026	2.3E-04	0.079 *	2.0E-04	0.21 *	0.026	0.16 *	0.044
2	2.6E-03	6.0E-05	0.025	2.0E-04	0.77 *	1.3E-04	0.39 *	8.1E-06
3	0.014	3.4E-04	0.081 *	2.0E-03	0.49 *	1.4E-05	0.98 **	6.6E-03
4	0.027	5.2E-04	0.048	4.5E-03	0.55 *	1.4E-04	0.85 **	9.8E-03
5	6.5E-05	7.6E-07	3.9E-03	1.2E-04	7.2E-03	3.2E-09	0.63 *	1.9E-05
6	4.6E-03	6.0E-05	9.0E-03	3.4E-04	0.32 *	3.8E-05	0.21 *	2.0E-03
7	1.5E-03	5.2E-05	0.067 *	0.019	0.16 *	1.4E-04	0.13 *	2.8E-02
8	0.022	1.4E-03	0.074 *	6.9E-03	0.21 *	6.6E-06	0.37 *	4.7E-03
9	9.1E-03	2.6E-04	0.023	2.1E-03	0.91 **	2.6E-06	0.55 **	9.0E-05
10	0.021	1.0E-03	8.1E-03	8.5E-04	0.15 *	3.1E-04	0.31 *	1.1E-04
Results from the datasets with <i>exponentially</i> distributed concentration								
no.	RMS WLS		MAD WLS		RMS OLS		MAD OLS	
	s_0^2	s_r^2	s_0^2	s_r^2	s_0^2	s_r^2	s_0^2	s_r^2
1	8.3E-04	5.2E-03	3.0E-03	0.019	6.1E-03	4.8E-05	0.027	5.9E-05
2	2.7E-05	3.6E-05	3.8E-04	3.4E-04	0.77 *	3.5E-09	0.12 *	4.7E-11
3	6.7E-05	6.2E-04	9.7E-04	4.8E-03	6.1E-03	8.3E-07	0.64 *	1.7E-06
4	1.5E-06	4.0E-06	5.3E-04	4.1E-04	3.8E-05	9.1E-10	0.049	7.8E-07
5	5.8E-04	3.9E-03	3.3E-04	0.010	1.5E-03	5.3E-08	1.4E-03	1.6E-06
6	1.9E-04	7.2E-04	5.1E-03	0.015	3.3E-03	3.1E-09	0.020	8.3E-06
7	2.3E-04	2.5E-03	3.3E-03	8.0E-03	1.4E-03	3.5E-05	0.20 *	8.9E-04
8	5.1E-05	4.0E-04	3.6E-04	3.2E-03	9.7E-03	7.5E-07	0.037	1.1E-07
9	1.2E-05	1.4E-04	2.5E-04	7.9E-03	0.040	1.9E-08	7.2E-04	2.0E-06
10	5.7E-04	3.9E-03	3.5E-03	7.2E-03	0.012	1.7E-05	0.025	1.1E-06

* indication of the estimates insignificant at $\alpha = 0.05$; ** indication of the insignificant and also negative estimates

When evaluating the regression procedures, the statistical significance of the s_0^2 and s_r^2 estimates, which represent the intercept and slope in Eq. (8), can also be assessed. The p-values of the t -test comparing the absolute values of these estimated squares with their standard errors are given in Table S12. Only the values of significant estimates, i.e., with $P \leq 0.05$, should be taken seriously. It means that the values of insignificant estimates should be considered indistinguishable from zero. Among the insignificant estimates, several negative s_0^2 estimates were found, which means that the corresponding s_0 values should be imaginary numbers. In such cases, the estimated s_0 value was given as zero (see Table S5 in ESM_1.pdf).

The results in Table S12 (ESM_1.pdf) show that all s_r^2 estimates found were significant. Insignificant estimates occurred only among the s_0^2 estimates, mainly within the estimates obtained from the uniformly distributed datasets. The RMS WLS procedure provided better estimates than the other regression procedures used. It was the only procedure that gave only significant estimates, mostly at $\alpha = 0.01$. When using the MAD or OLS method instead of the RMS or WLS method, the quality of the estimates was worse, the number of the estimates that were significant only at $\alpha = 0.05$ increased, and insignificant or even negative s_0^2 estimates appeared. From the uniformly distributed datasets, the MAD OLS, RMS OLS and MAD WLS procedures provided 10, 9 and 4 insignificant estimates, respectively. In the first and second cases, three and one estimates were even negative. The MAD OLS procedure turned to be the worst regression procedure.

Text S7 Correlation between the parameter estimates obtained by all used procedures

Tables S13 and S14 (ESM_1.pdf) summarise Spearman's correlation coefficients calculated from all possible pairs of the sets of s_0 or s_r estimates obtained by all procedures used and by their variants from the examined datasets distributed uniformly and exponentially. The highest numbers of significant correlations were found for the s_r and s_0 estimates obtained from the datasets with the uniform and exponential distribution, 17 out of 28 and 12 out of 28, respectively. These cases mainly included the significant correlations between the estimates obtained by the variants of the proposed procedure and RMS WLS procedure, which have been mentioned in Subchapter "Comparing the estimates obtained by the newly proposed and regression procedures". It should be recalled that in the datasets with the uniform and exponential distribution the differences suitable for estimating s_r and s_0 by the proposed procedure prevailed, respectively. The estimate sets obtained by the RMS WLS and RMS OLS procedures were correlated with each other only in the two above cases of combination of parameter and distribution type. Interestingly, the s_r estimate sets obtained by the four variants of the proposed procedure from the datasets with the uniform distribution correlated with the estimate sets obtained by the RMS OLS and MAD OLS procedure in all four and three cases, respectively. The estimates obtained by the MAD WLS and MAD OLD procedures were correlated with each other only for s_0 or s_r estimates from exponentially distributed datasets. Of course, all significant correlations found were positive. No correlation was found between the estimate sets obtained by the proposed and MAD WLS procedure.

Text S8 References

- 1s. Molin P, Abdi H (1998) New Tables and numerical approximation for the Kolmogorov-Smirnov/Lilliefors/Van Soest test of normality. Technical report, University of Bourgogne. <http://www.utd.edu/~herve/MolinAbdi1998-LillieforsTechReport.pdf>. Accessed 10 Oct 2021
- 2s. Brown S (2008) Statistics: Measures of Shape: Skewness and Kurtosis. <https://brownmath.com/stat/index.htm>. Accessed 10 Oct 2021

Table S13. Matrices of Spearman's correlation coefficients evaluating the relationship between the parameter estimates obtained by using the various procedures from the datasets with $n = 200$ distributed **uniformly**; the coefficients for the s_0 and s_r estimates are in the right upper and left lower part, respectively

Procedure	Proposed $c = 2.6$	Proposed $c = 4$	Proposed 85/150	Proposed 100/100	RMS WLS	MAD WLS	RMS OLS	MAD OLS
Proposed $c = 2.6$	s_r s_0	-0.042	-0.115	0.042	0.842 *	0.576	-0.091	0.202
Proposed $c = 4$	0.588	s_r s_0	0.782 *	0.552	-0.309	0.103	0.430	0.202
Proposed 85/150	0.673 *	0.879 *	s_r s_0	0.721 *	-0.115	0.224	0.345	0.362
Proposed 100/100	0.515	0.842 *	0.903 *	s_r s_0	-0.018	-0.115	0.139	0.423
RMS WLS	0.709 *	0.879 *	0.782 *	0.721 *	s_r s_0	0.564	-0.406	0.018
MAD WLS	-0.030	0.309	0.370	0.333	0.406	s_r s_0	0.030	0.117
RMS OLS	0.794 *	0.770 *	0.855 *	0.855 *	0.770 *	0.321	s_r s_0	0.693 *
MAD OLS	0.806 *	0.491	0.685 *	0.636 *	0.576	0.309	0.915 *	s_r s_0

* indication of the coefficients significant at $\alpha = 0.05$, i.e., > 0.648 (see [26])

Table S14. Matrices of Spearman's correlation coefficients evaluating the relationship between the parameter estimates obtained by using the various procedures from the datasets with $n = 200$ distributed *exponentially*; the coefficients for the s_0 and s_r estimates are in the right upper and left lower part, respectively

Procedure	Proposed 110/90	Proposed 100/100	Proposed 100/120	Proposed 100/140	RMS WLS	MAD WLS	RMS OLS	MAD OLS
Proposed 110/90	s_r s_0	0.915 *	0.806 *	0.939 *	0.709 *	0.430	0.564	0.067
Proposed 100/100	0.855 *	s_r s_0	0.939 *	0.964 *	0.721 *	0.467	0.636	0.139
Proposed 100/120	0.418	0.685 *	s_r s_0	0.939 *	0.758 *	0.370	0.636	0.103
Proposed 100/140	0.600	0.564	0.539	s_r s_0	0.794 *	0.503	0.636	0.188
RMS WLS	0.200	0.188	0.321	0.455	s_r s_0	0.539	0.915 *	0.370
MAD WLS	-0.006	-0.030	0.067	0.455	0.273	s_r s_0	0.479	0.855 *
RMS OLS	-0.212	-0.285	-0.164	0.055	0.624	0.418	s_r s_0	0.333
MAD OLS	-0.394	-0.491	0.018	0.285	0.333	0.733 *	0.479	s_r s_0

* indication of the coefficients significant at $\alpha = 0.05$, i.e., > 0.648 (see [26])